

Cloud Computing

Virtual Resources and Distributed Cloud Applications

Chapter 4: Scalable and Fault-Tolerant Applications



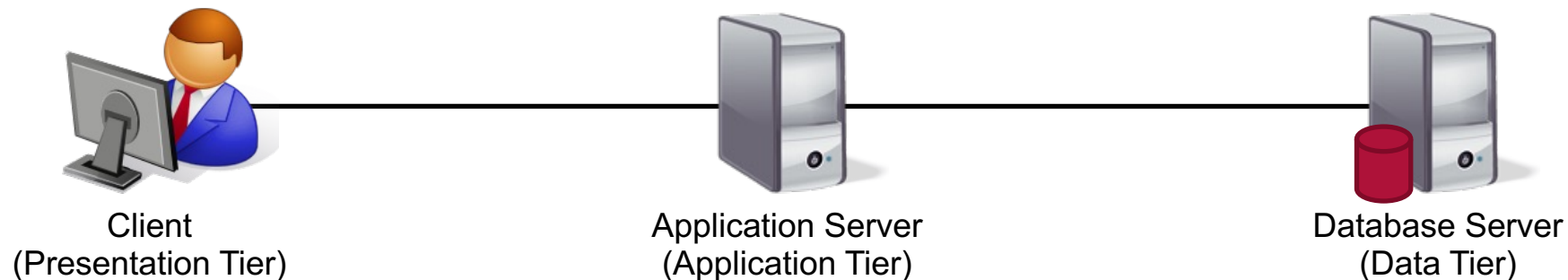
Distributed and Operating Systems
Electrical Engineering and Computer Science
Technische Universität Berlin

Overview

- **Intro**
- Scaling Stateless Components
 - Kubernetes Auto-Scaling
- Partitioning
 - Amazon DynamoDB
- Replication and Consistency
- CAP Theorem

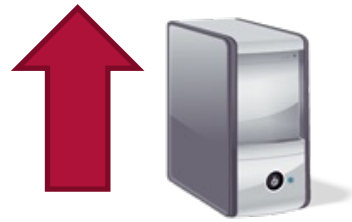
How to Achieve Scalability / Availability?

- Let's start with a classic 3-tier architecture:
 - What happens if we increase the number of clients?
 - What happens if one of the components goes down?



- What are techniques to achieve scalability/availability?

How to Achieve Scalability / Availability?



Vertical Scalability (Scaling Up)

Idea	Increase performance of a single node (more cores, memory, ...)
Pro:	Good speedup up to a particular point
Con:	Beyond that point, speedup becomes very expensive



Horizontal Scalability (Scaling Out)

Idea	Increase number of nodes
Pro:	Cheap to grow total amount of resources
Con:	More difficult for applications to efficiently utilize the resources

Commodity Servers and the Cloud

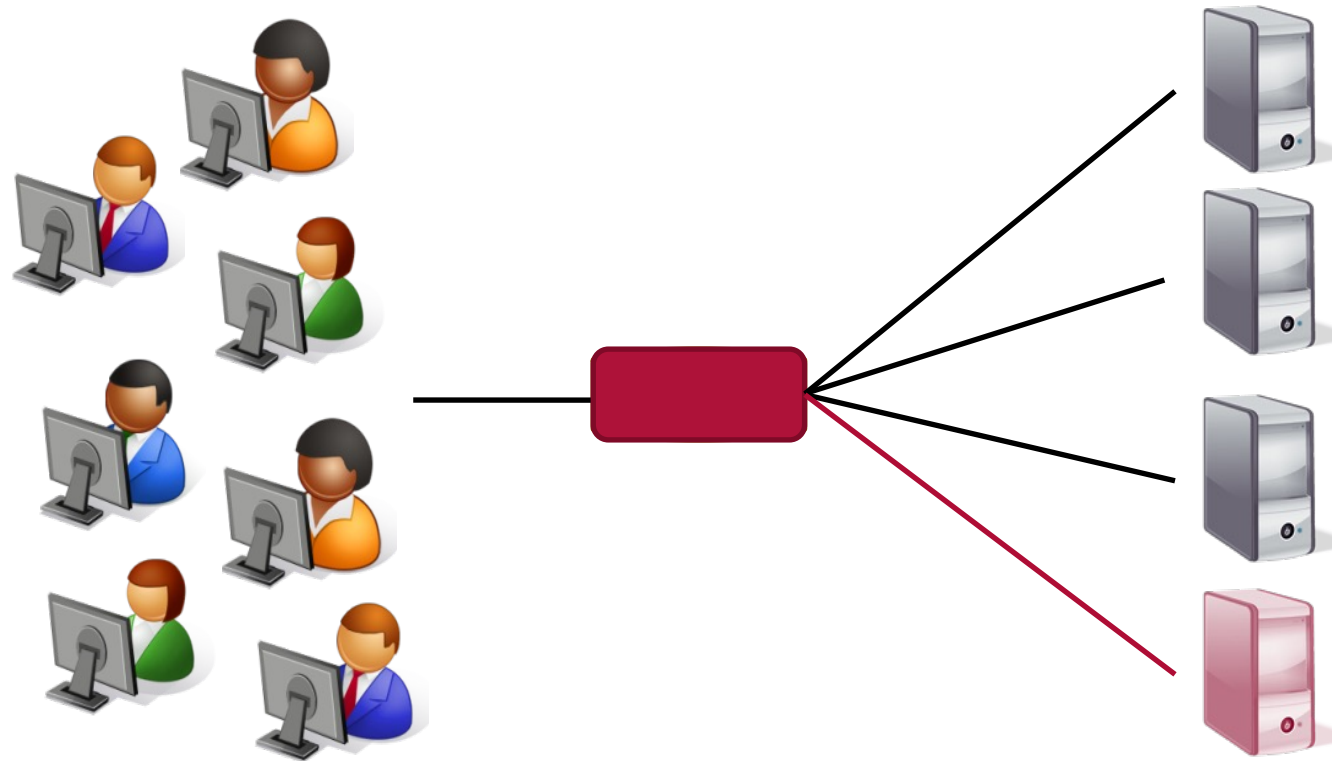
- Hardware trend: CPUs do not become faster as fast as they used to
→ datacenters today house large sets of inexpensive commodity servers that fail eventually
- Scaling applications therefore typically means scaling-out across many nodes
- VMs and containers make it easy to replicate services across many nodes, especially with IaC and automation/orchestration tools

The Fallacies of Distributed Computing

1. The network is reliable.
2. Latency is zero.
3. Bandwidth is infinite.
4. The network is secure.
5. Topology does not change.
6. There is one administrator.
7. Transport cost is zero.
8. The network is homogeneous.

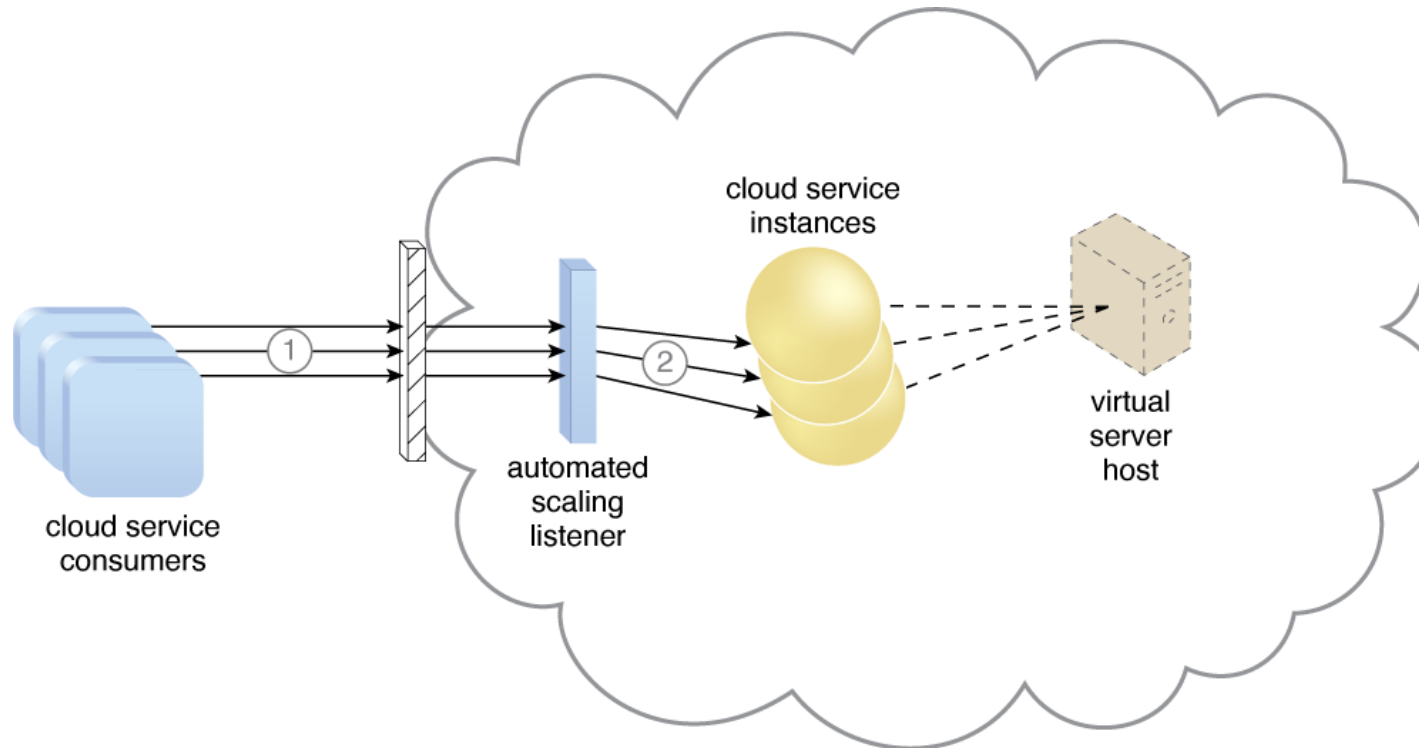
Dynamic Scaling Architectures

- Idea: pre-defined scaling conditions that trigger the dynamic allocation of resources from shared pools



Simple Dynamic Scaling Architecture (1/3)

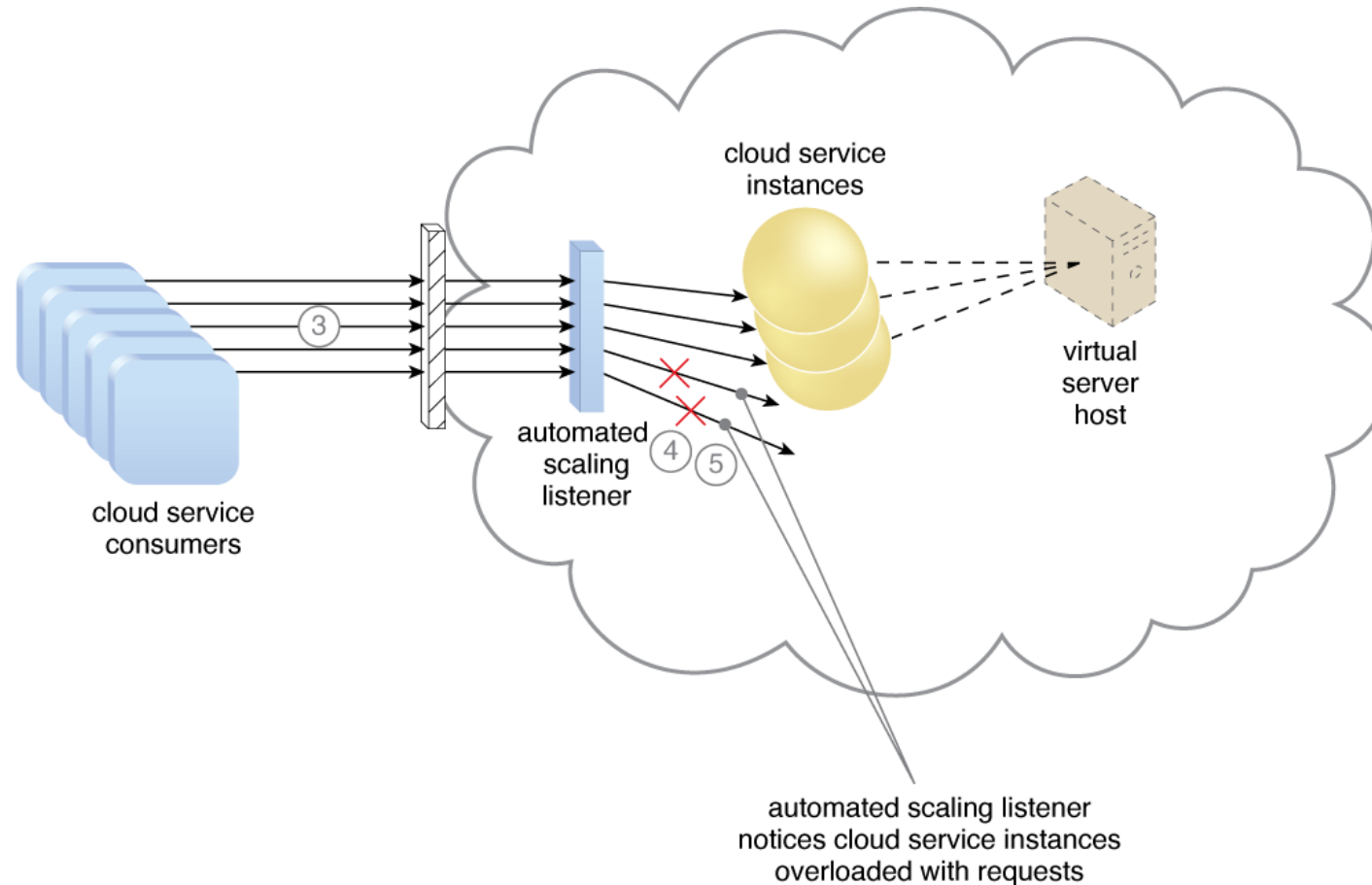
- Cloud consumers sending requests



Copyright © Arcitura Education

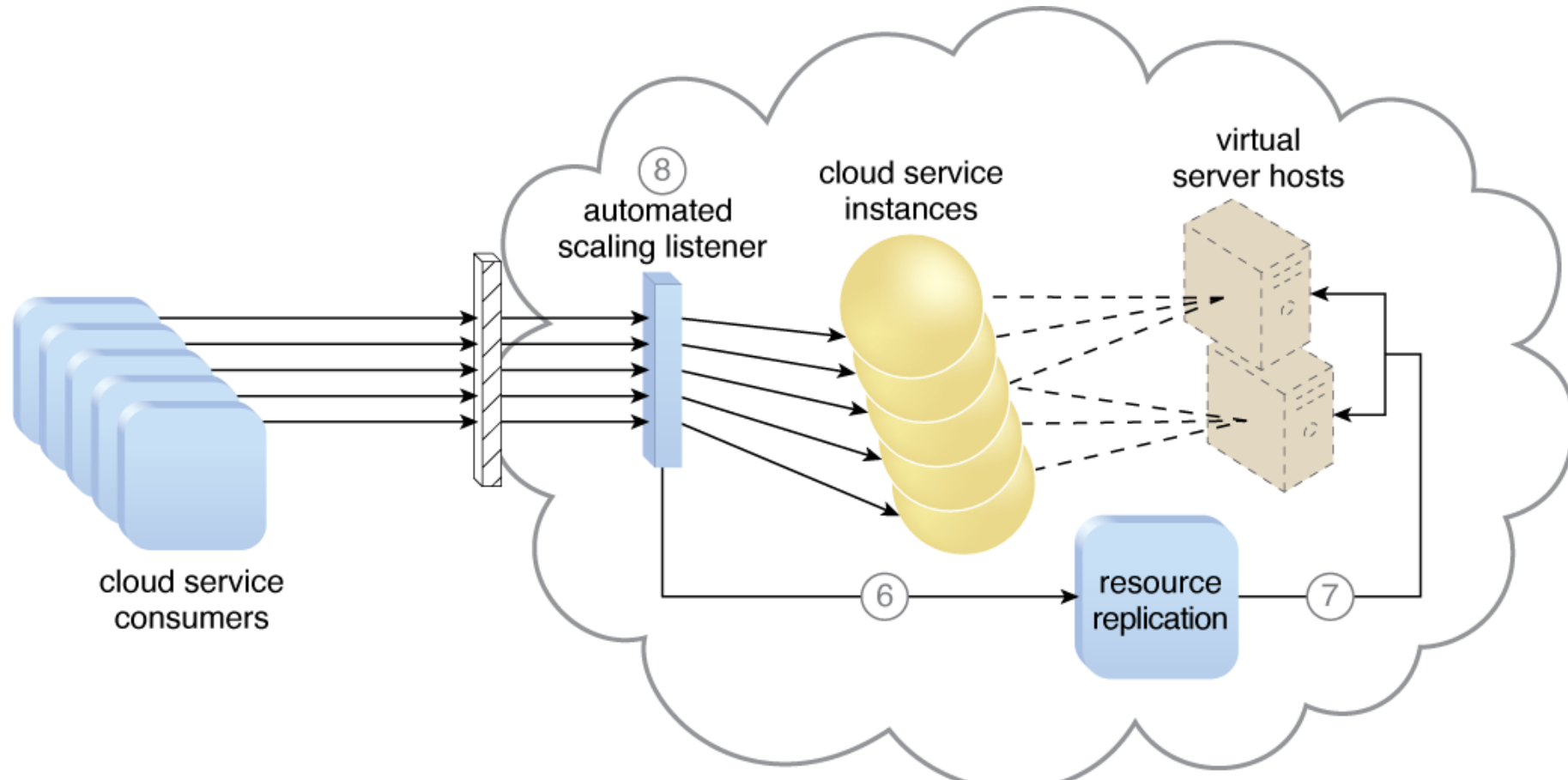
Simple Dynamic Scaling Architecture (2/3)

- Requests overload the existing resources, timeouts exceeded



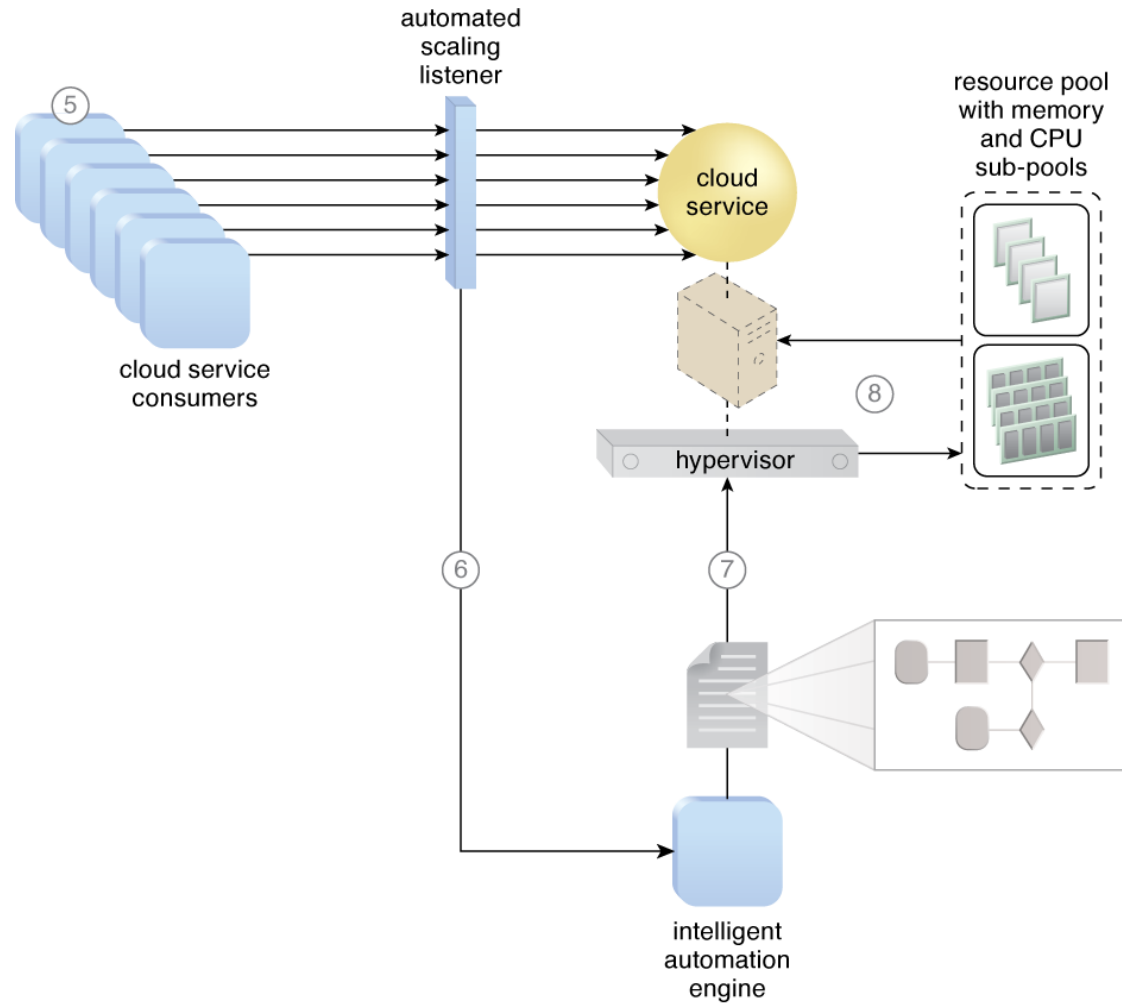
Simple Dynamic Scaling Architecture (3/3)

- Scaling up by deploying additional resources



More Advanced Elasticity Architecture

- Scaling up by deploying more resources using an *intelligent automation engine*



How to Achieve Scalability?

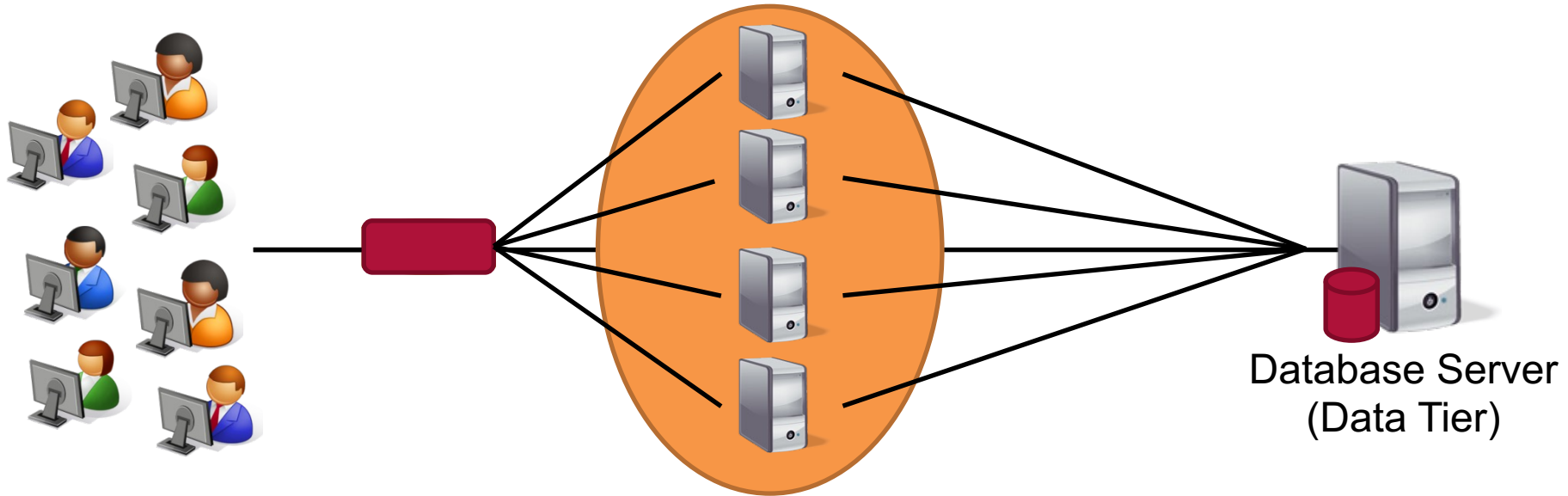
- Multi-tier application scalability entails two questions:
 - Where is the bottleneck in the architecture?
 - Is the bottleneck component stateless or stateful?
- Stateless components
 - Component maintains no internal state beyond a request
 - Examples: DNS server, web server with static data, ...
- Stateful components
 - Component maintains state beyond request (state is required to process the next request)
 - Examples: mail server, stateful web server, DBMS, ...

Overview

- Intro
- **Scaling Stateless Components**
 - Kubernetes Auto-Scaling
- Partitioning
 - Amazon DynamoDB
- Replication and Consistency
- CAP Theorem

Scalability with Stateless Components

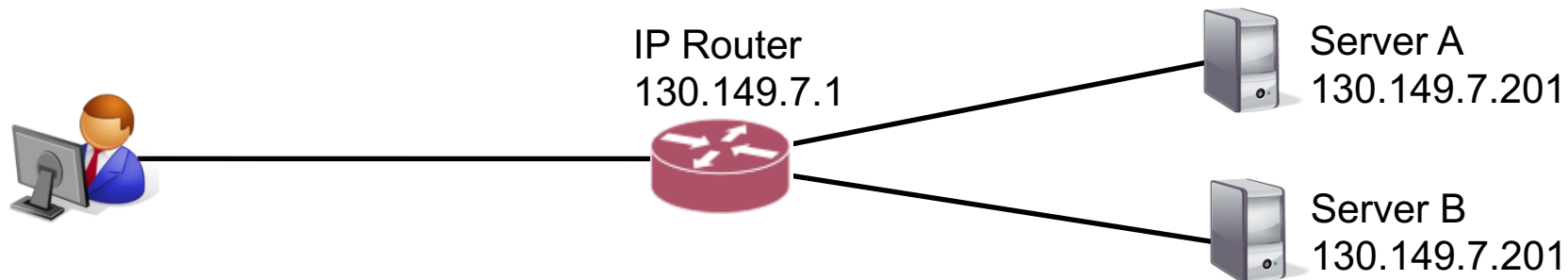
- Let's assume stateless servers are the bottleneck → more instances of the stateless component:



- How do clients figure out which server to contact?
 - Clients know list of servers (e.g. P2P servers)
 - Introduction of a load balancer

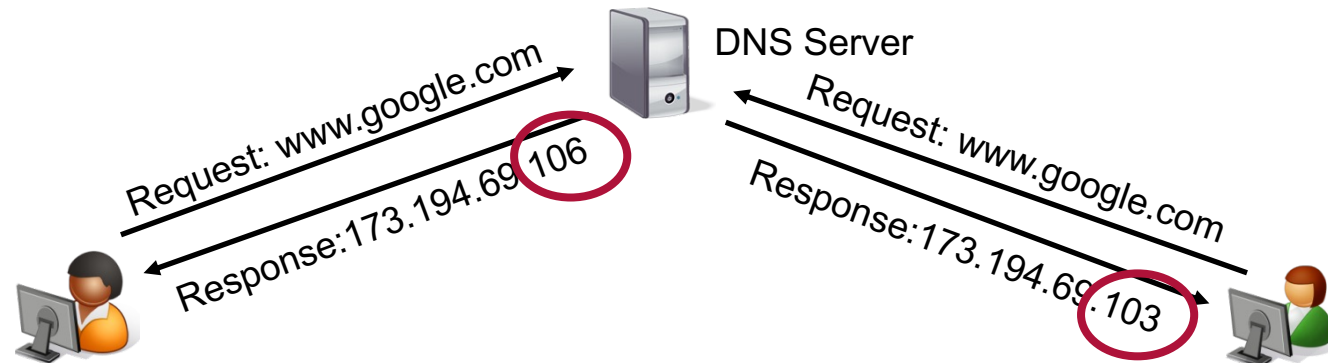
Possible Levels of Stateless Load Balancing (1/3)

- Load balancing on IP level
 - Balancing is implemented by IP routers
 - ◆ Multiple devices share one IP address (IP anycast)
 - ◆ Routers route packet to different locations
 - Requirements for applicability
 - ◆ Request must fit in one IP packet
 - ◆ Control over routers
 - Examples: DNS root server (mostly for availability)



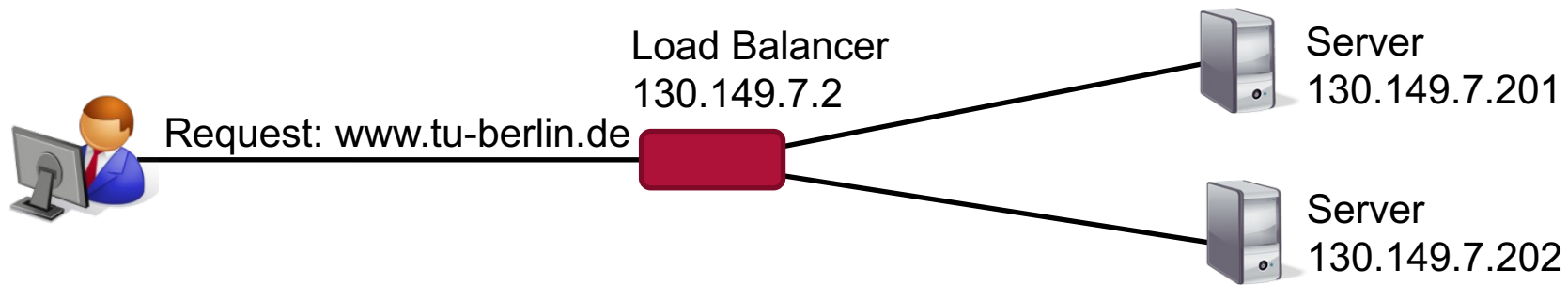
Possible Levels of Stateless Load Balancing (2/3)

- Load balancing on DNS level
 - Balancing is implemented by DNS servers
 - ◆ DNS servers resolve domain names to IP addresses
 - Requirements for applicability
 - ◆ Control over DNS server(s)
 - ◆ Stable load characteristics (think of DNS caching)
 - Examples: Various big websites, e.g. *google.com*



Possible Levels of Stateless Load Balancing (3/3)

- Load balancing by distinct load balancer
 - Deliberately distributes requests among machines
 - Clients first send request to load balancer, then to target
- Requirements for applicability
 - No network bottleneck
- Examples: Various websites, Amazon Elastic LB



Strategies for Stateless Load Balancing

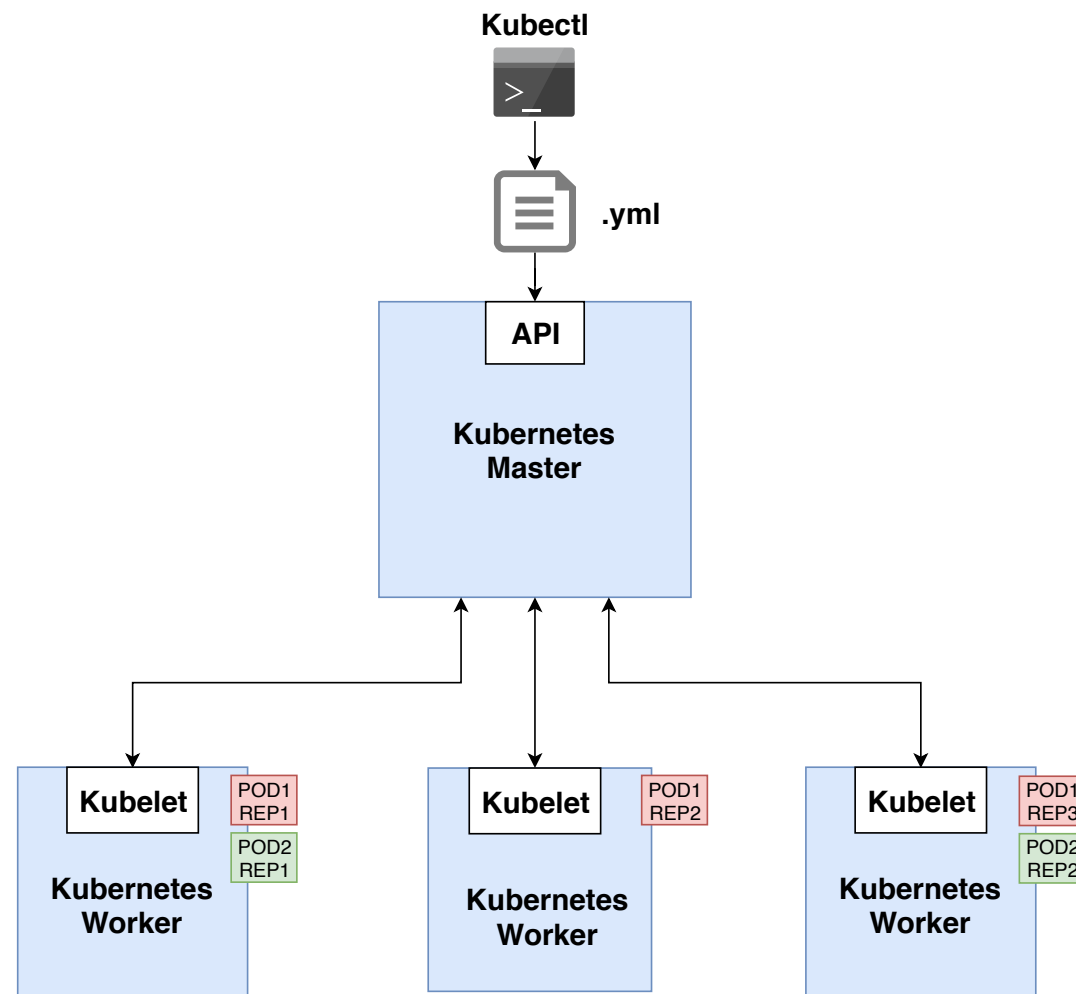
- Load balancing on different levels can be combined
 - For example, distribute network load among distinct load balancers first with DNS load balancing
- Different strategies for the actual load balancing
 1. Round robin LB
 - ◆ Simple, good if all request cause roughly the same load
 2. Feedback-based LB
 - ◆ Servers report actual load back to load balancer
 3. Client-based LB
 - ◆ Choose server with smallest network latency for client

Overview

- Intro
- Scaling Stateless Components
 - **Kubernetes Auto-Scaling**
- Partitioning
 - Amazon DynamoDB
- Replication and Consistency
- CAP Theorem

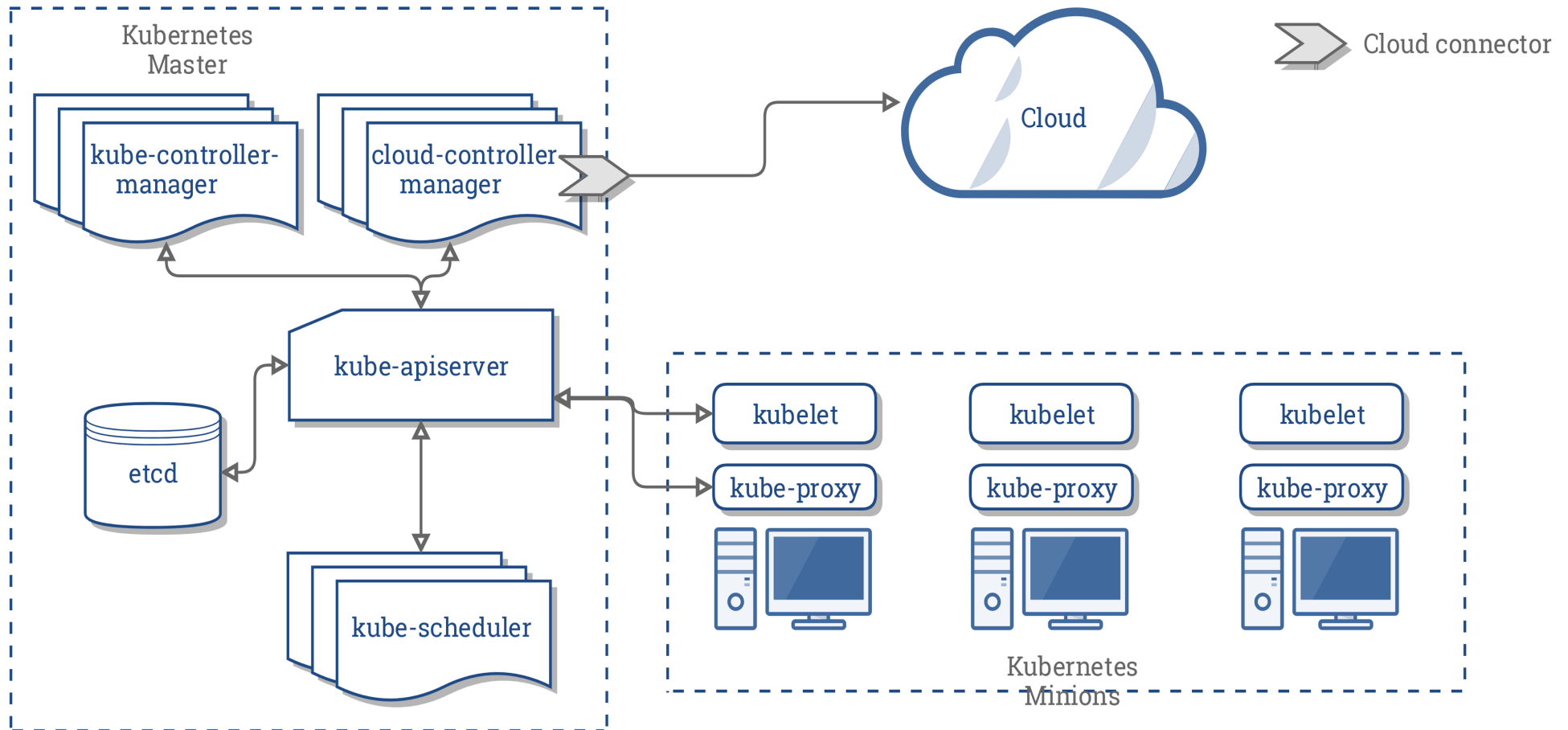
Recap: Kubernetes Cluster

- Manages collection of nodes (either bare-metal or virtual machines)
- Runs sets of replicated containers (called *Pods*)



Kubernetes

Horizontal Pod Autoscaler



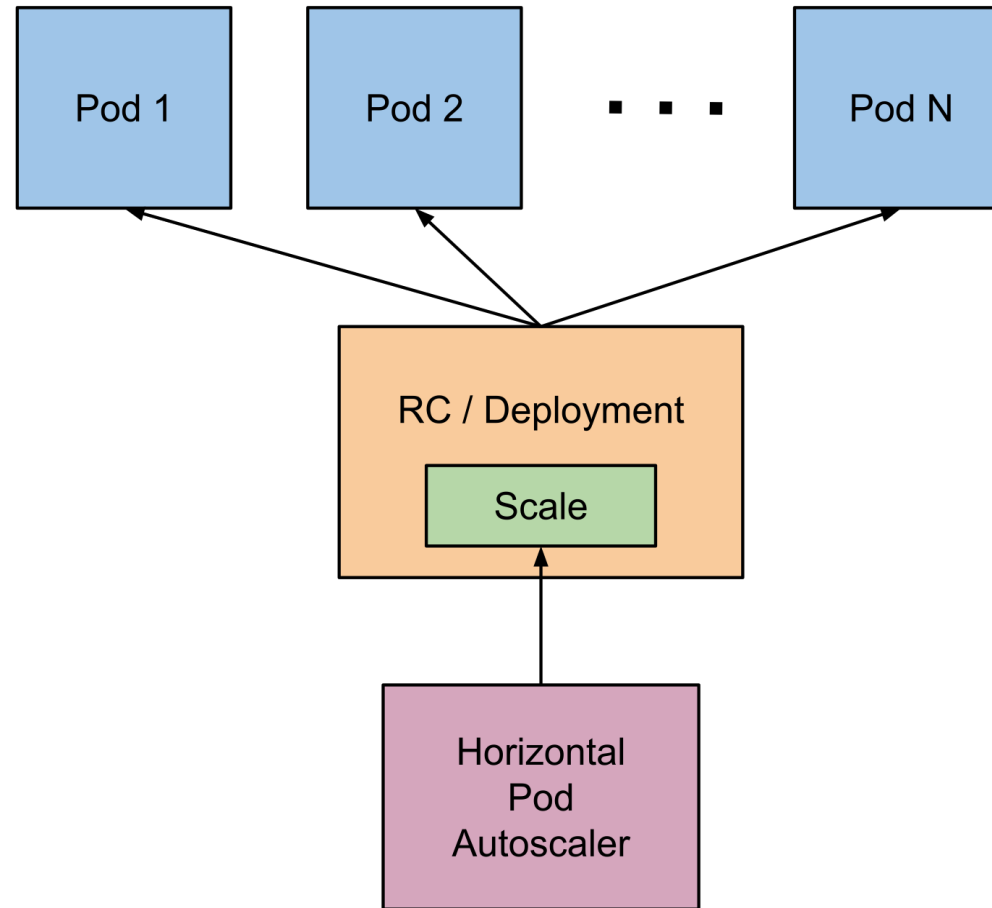
Kubernetes

Horizontal Pod Autoscaler

- Automatically scales number of pods in a replication controller
- Based on
 - CPU utilization or custom metrics
 - Target value for the metric
- Autoscaler checks metrics every 30 seconds
- Sets the number of replications to optimize the metric towards the target value
 - By creating more pods
 - or by reducing the number of pods

Kubernetes

Horizontal Pod Autoscaler



Kubernetes

Horizontal Pod Autoscaler

- Number of replicas is scaled by the quota of average current metric value and target value
- No scaling if this quota is within a 0.1 tolerance

```
desiredReplicas =  
    [ currentReplicas * (currentMetricValue/desiredMetricValue) ]
```

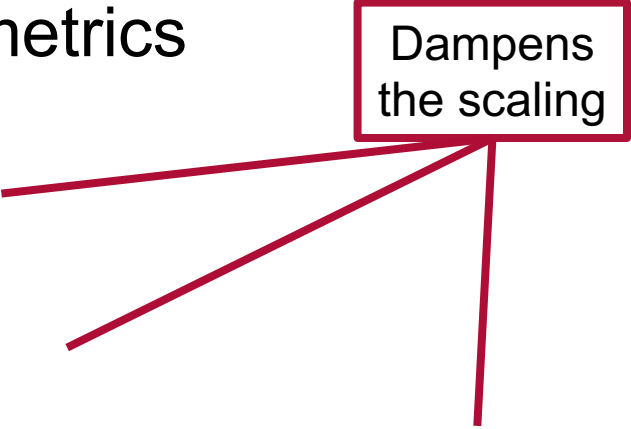
→ This assumes linear scaling!

- Autoscaler scales to the highest number of desired replicas in a sliding window of 5 minutes
 - Quick responses to more load, but reduces *thrashing*

Sometimes Metrics are not Available

- Discard pods that are currently being shut down
- Normally running pods with missing metrics
 - If result without these would be to scale up: assume metric to be 0
 - If result without these would be to scale down: assume metric to be 1
- Assume metric to be 0 for pods that are not yet ready

Dampens
the scaling

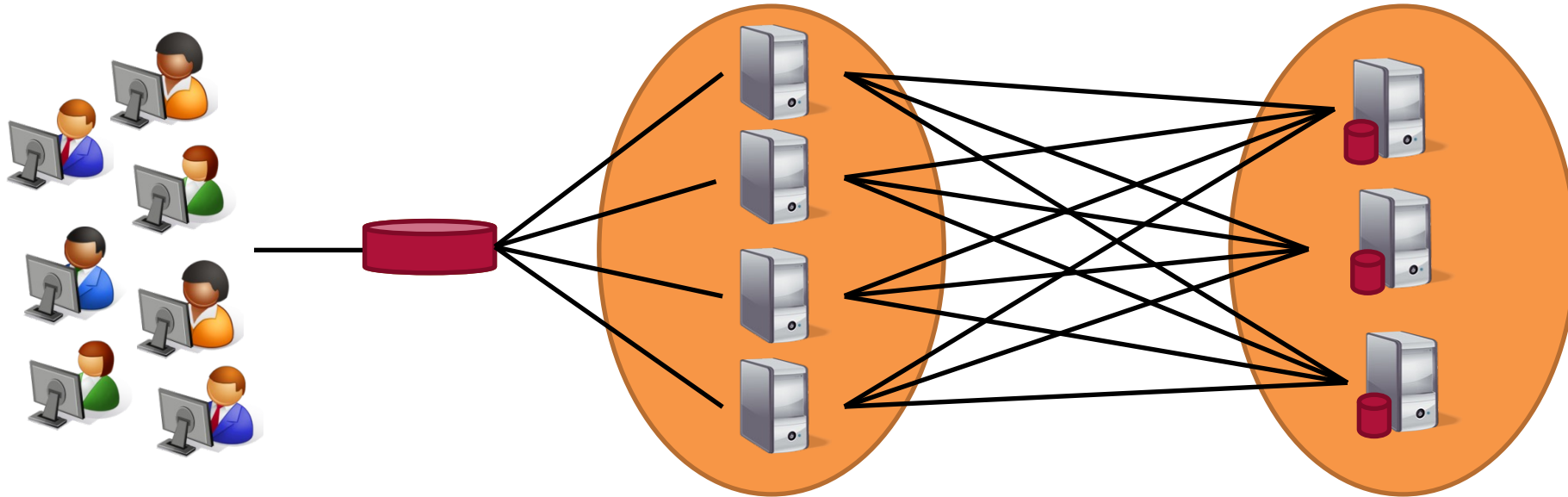


Overview

- Intro
- Scaling Stateless Components
 - Kubernetes Auto-Scaling
- **Partitioning**
 - Amazon DynamoDB
- Replication and Consistency
- CAP Theorem

Scalability with Stateful Components

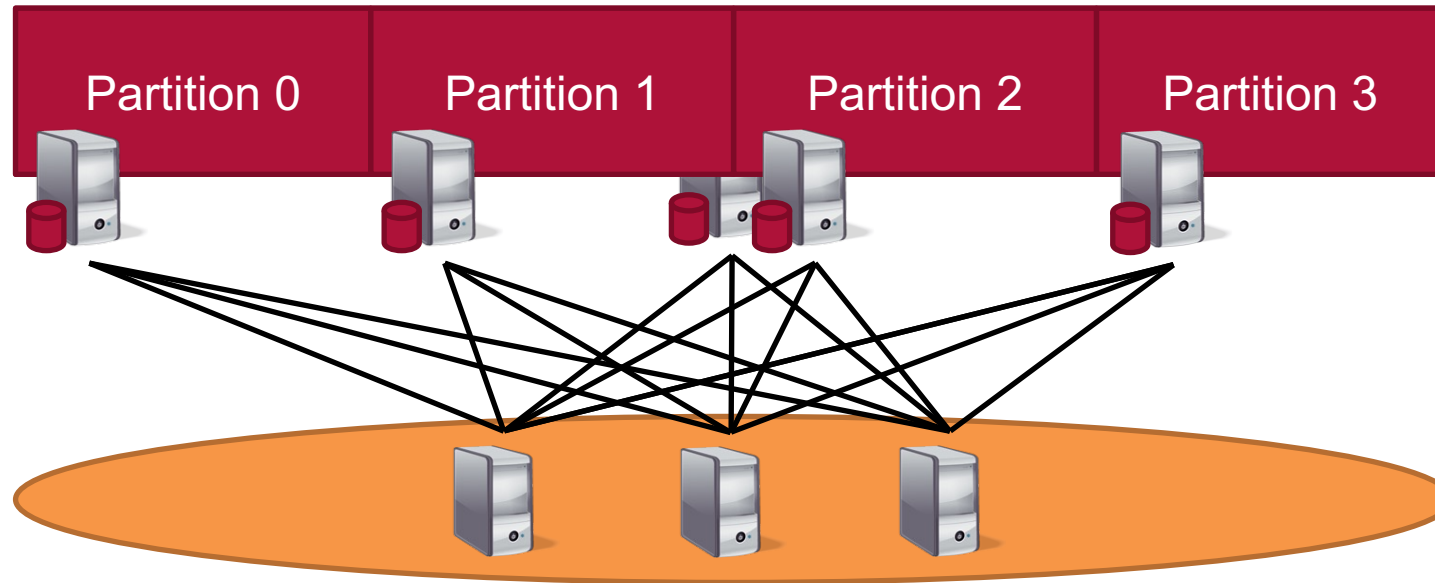
- Let's assume the data tier is the bottleneck now



- Database is stateful (stores data beyond requests) → requests from the same client must be handled by the same instance of the database server

Scalability with Stateful Components: Partitioning

- Idea: Divide data into distinct, independent part: Each server is responsible for one or more parts



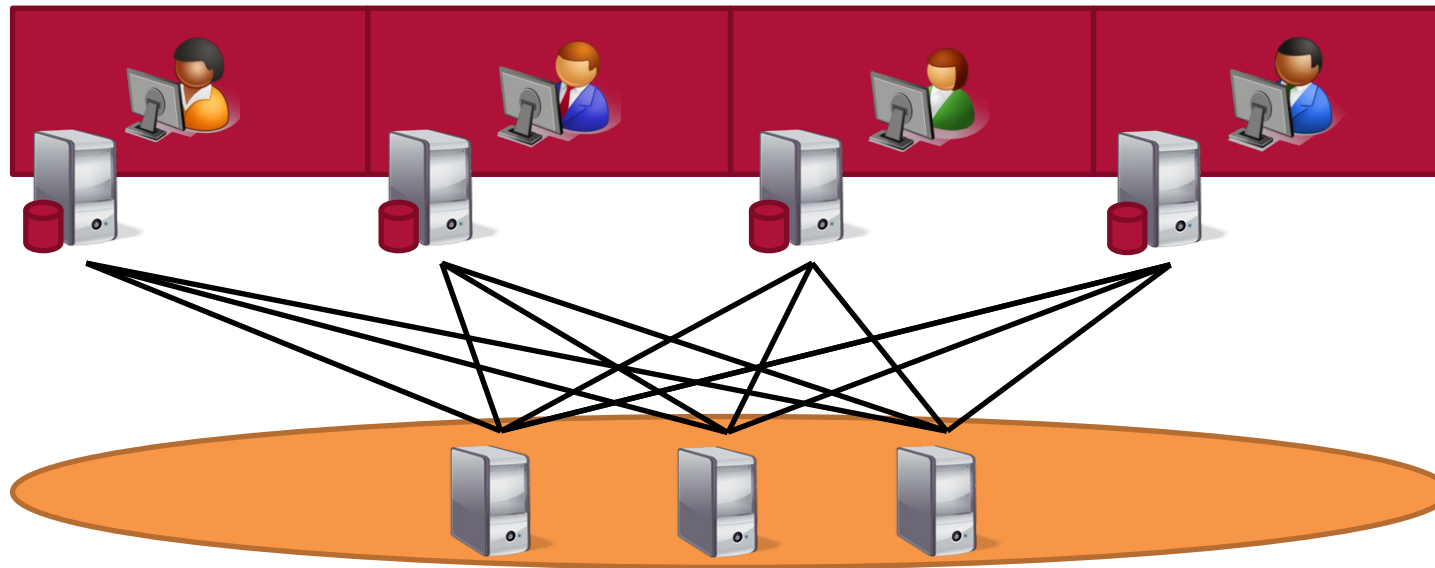
- Partitioning improves just scalability, not availability: Each data item is only stored in one partition!

How to Partition the Data?

- General consideration for partitioning and scalability
 - Data is now spread across several machines
 - Particular tasks might require to pull data from multiple servers now → causes additional network traffic
- Network is a scarce resource with limited scalability
 - Becomes scarcer the more machines you add
- For good scalability, the goal of every partitioning scheme is typically to reduce network communication
 - However, this is highly application-specific...

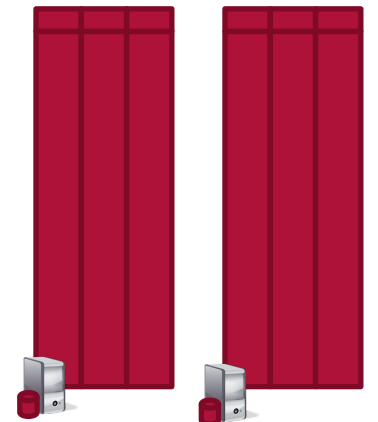
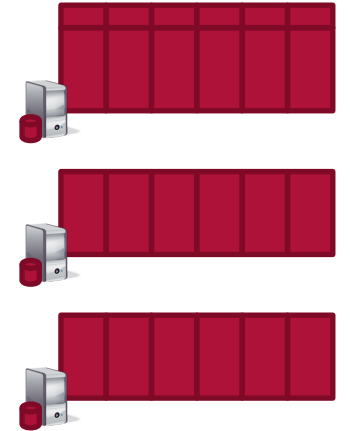
Popular Partitioning Schemes (1/2)

- Partitioning per tenant
 - Put different tenants on different machines
 - Pro: In clouds, tenants are expected to be isolated
 - No network traffic between machines, good scalability
 - Con: Tenants cannot scale beyond one machine



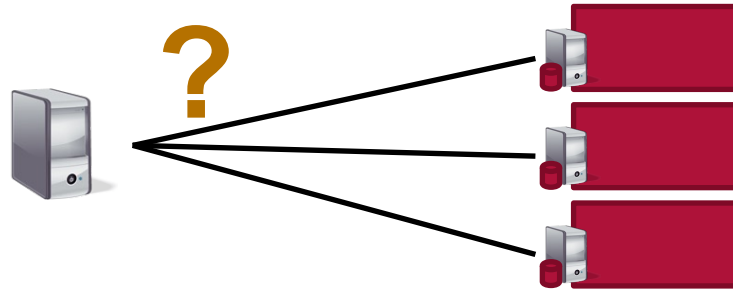
Popular Partitioning Schemes (2/2)

- Horizontal partitioning (relational databases)
 - Split table by rows
 - Put different rows on different machines
 - Reduced number of rows, reduced indices
 - Done by Google BigTable and MongoDB
- Vertical partitioning (relational databases)
 - Split table by columns
 - Not very common to improve scalability (?)



How to Distribute Data Among Partitions?

- Define key attribute $k \in K$ on the data item to be stored
- Define a globally-known partition function p so that



$$p : K \rightarrow P$$

P is set of available partitions

- Characteristics of p have significant impact on
 - Load balancing
 - Scalability

Classes of Partition Functions

- Hash partitioning
 - Desired property: Uniform distribution
 - Pro: Good load balancing characteristics
 - Con: Inefficient for range queries, typically requires data reorganization when number of partitions changes
- Range partitioning
 - Desired property: If $k_1, k_2 \in K$ are close, $p(k_1), p(k_2) \in P$ shall also be close to each other
 - Pro: Efficient for range queries and scaling individual partitions
 - Con: Often poor load balancing properties

Overview

- Intro
- Scaling Stateless Components
 - Kubernetes Auto-Scaling
- Partitioning
 - **Amazon DynamoDB**
- Replication and Consistency
- CAP Theorem

Amazon DynamoDB^[4]

- Highly-available key-value store, provided as a managed service by Amazon
- Primary design goals
 - High scalability
 - ◆ e.g. 1000s of servers
 - High availability
 - ◆ Specifically, support for “always write”
 - High performance
 - ◆ Particularly, small latency (single digit milliseconds) and high throughput (20 million requests per second)
- Sacrifices consistency to achieve these goals

Amazon DynamoDB: Design Principles

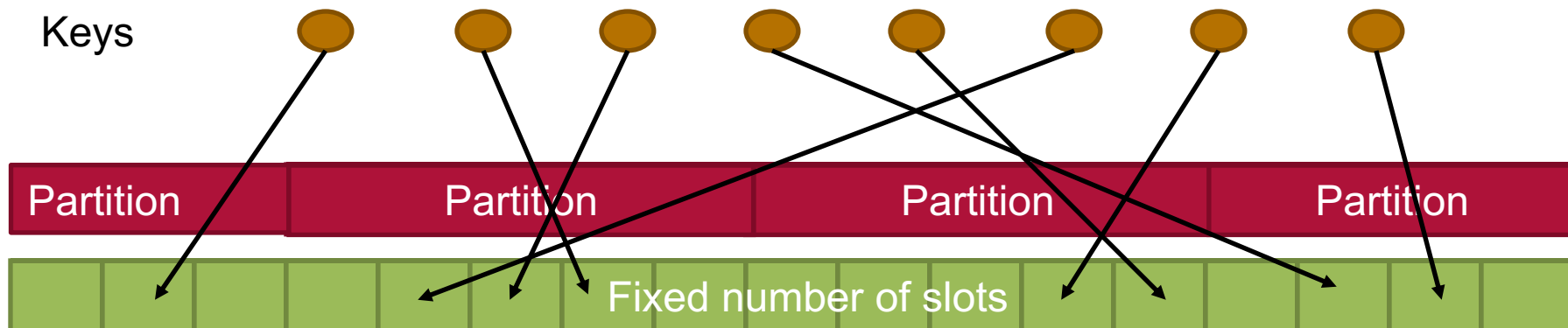
- Dynamo follows peer-to-peer approach
 - No server is more important than any other
 - No single point of failure
- Nodes can be added/removed incrementally at runtime
 - Weak consistency guarantee (eventual consistency)
 - In terms of the CAP theorem, DynamoDB is AP
- No hostile environment, all servers follow the rules
 - No security issues must be considered

How to Partition Data Among the Servers?

- Dynamo designed to be a key-value store
 - ➔ No support for range queries needed, so no need for range partitioning
- Hash partitioning has good load balancing properties
 - But how to avoid data reorganization when servers can be added and removed dynamically?
- Solution: Hash function no longer maps to partitions, but to a fixed number of slots
 - Variant of classic hashing called *consistent hashing*

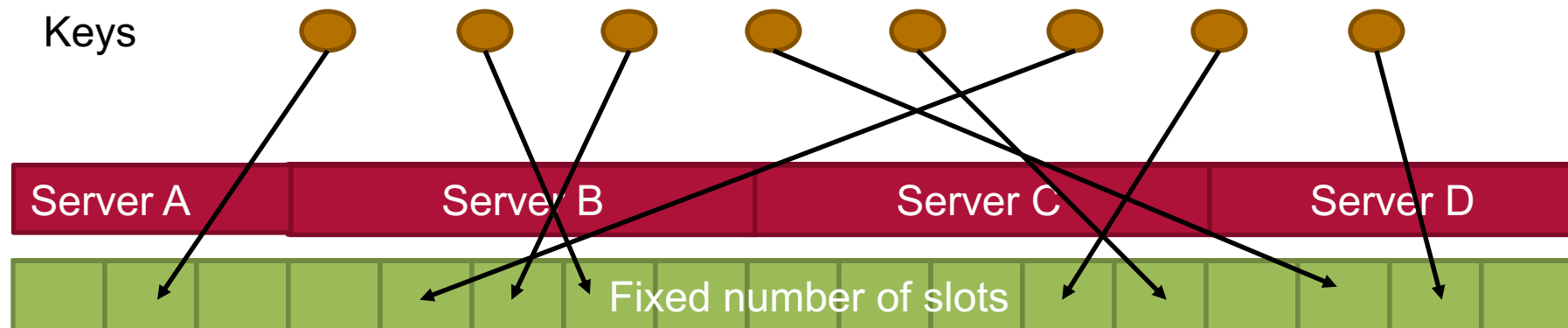
Consistent Hashing_[6]

- Idea: Additional mapping between slots and partitions
 - Hash function maps keys to large but fixed number of slots (for example 2^{128} slots)
 - A partition now covers a consecutive range of slots
 - ◆ A server can be in charge of one or more partitions
 - ◆ Size of a partition is variable



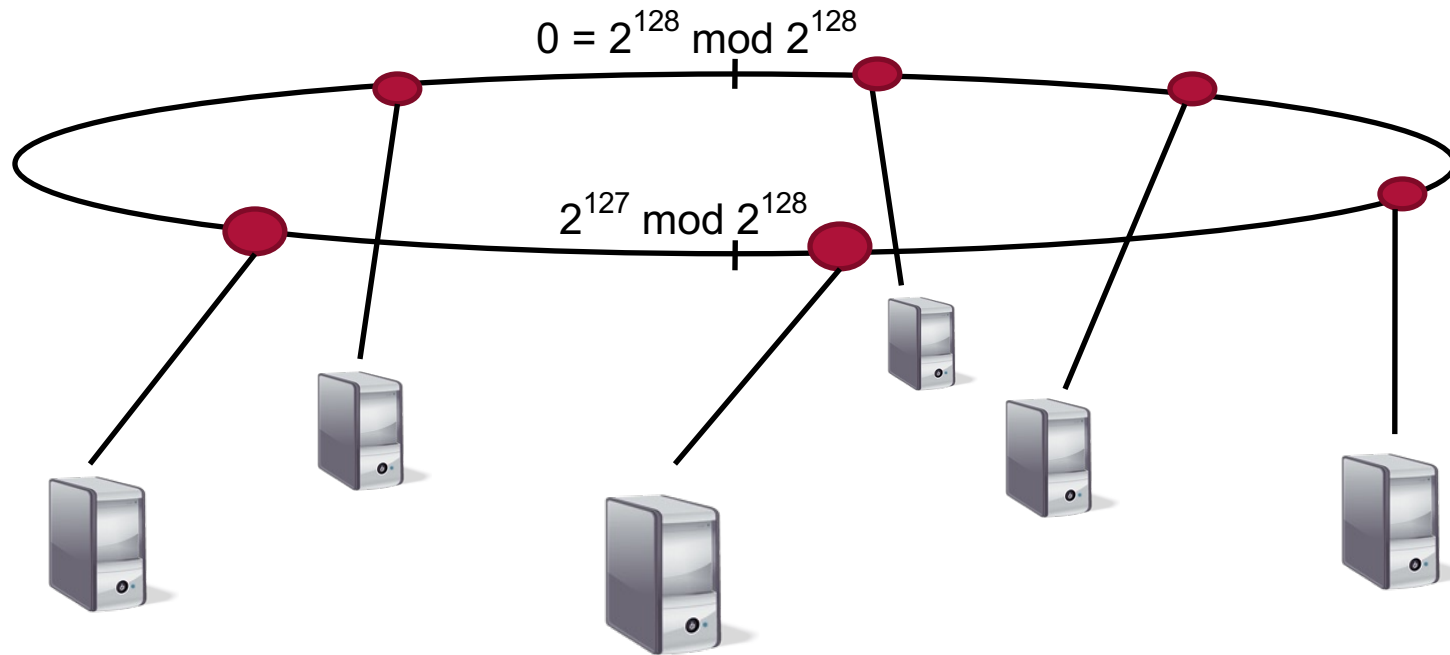
Distributed Hash Tables

- Consistent hashing is the basis for distributed hash tables (DHTs)
 - Each server takes at least one partition → each server is responsible for a continuous range of slots
 - Upon arrival/departure of server only $O(\#Keys/\#Servers)$ data items must be reorganized



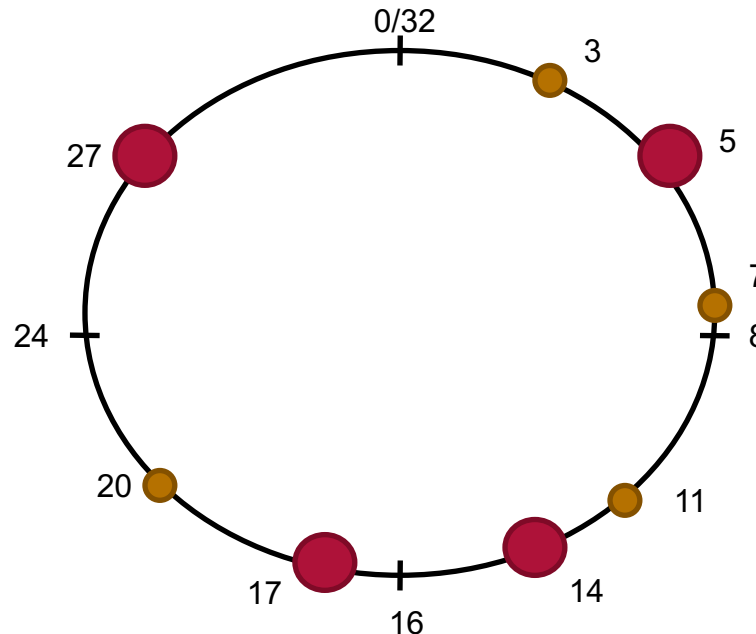
Amazon DynamoDB's Partitioning Algorithm (1/2)

- Slots form a circular ID space
 - All servers hash their ID with an MD5 function
 - Hashing result determines server's position on the ring



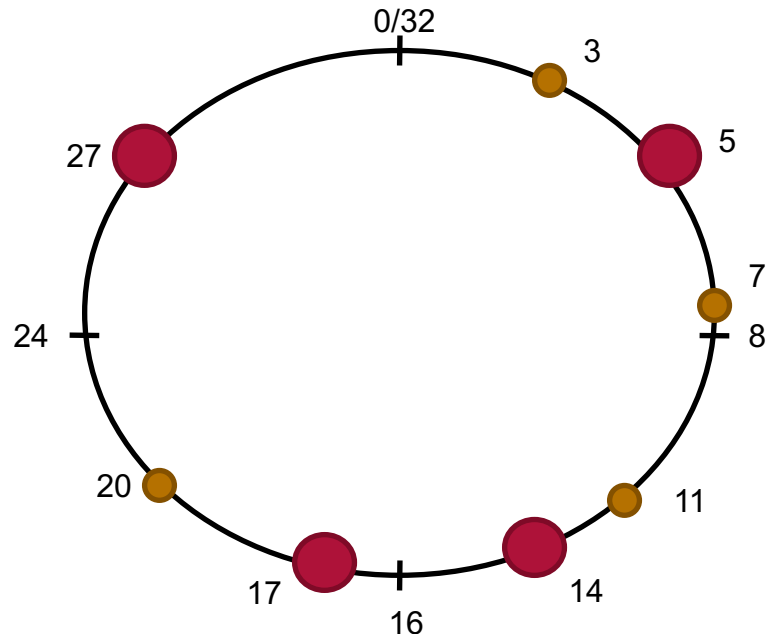
Amazon DynamoDB's Partitioning Algorithm (2/2)

- To distribute data, key of data also hashed with MD5
 - Result of hashing is a position in the circular ID space
 - Rule: Server is responsible for all preceding IDs (i.e. slots) up to and including its own ID



Routing in Amazon DynamoDB

- Each server maintains full routing table (ID to IP)
 - Each server can determine directly where a data item is based on the mapping rule and routing table!
 - One hop routing keeps latencies small

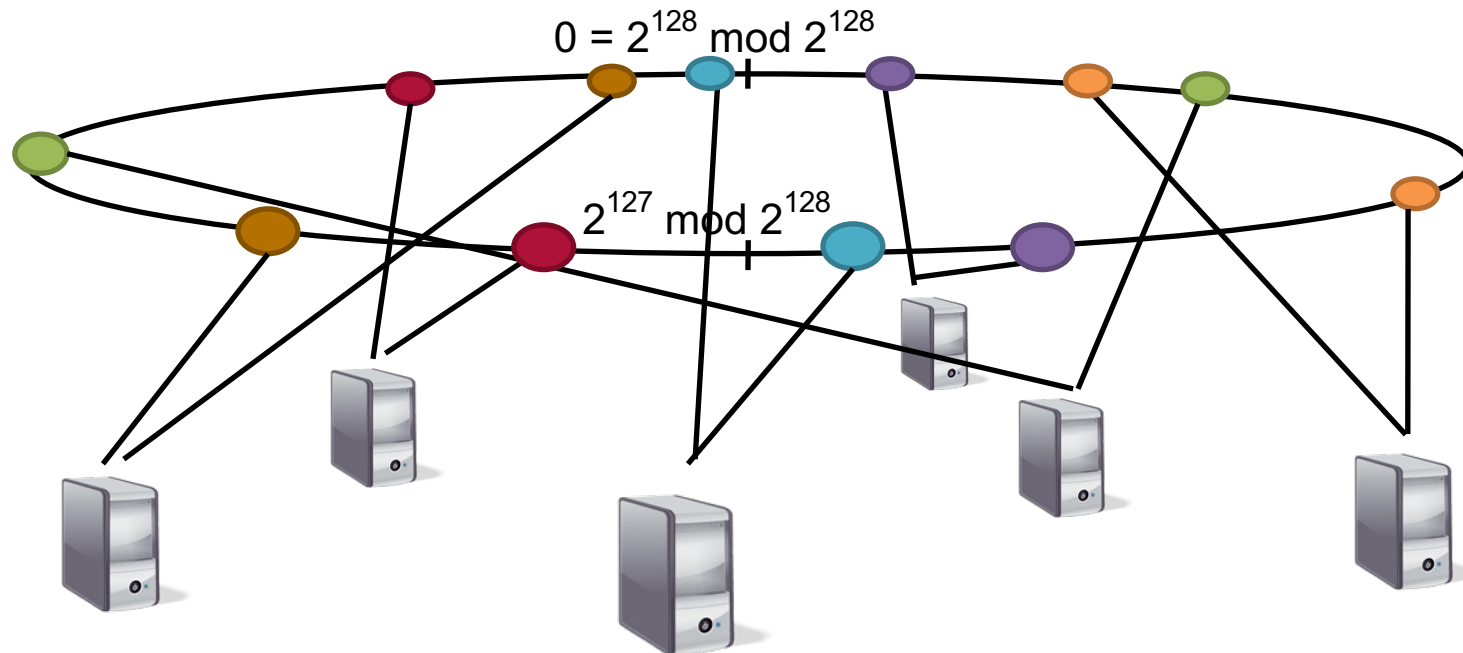


Example routing table in DynamoDB

ID	IP Address
5	192.168.1.2
14	192.168.1.18
17	192.168.1.115
27	192.168.1.98

Virtual Servers for Load Balancing

- Despite uniform distribution over ID space, the servers may receive skewed number of requests
- Solution: Each server appears on multiple positions on the ring (known as virtual servers)



Replication in Amazon DynamoDB

- Servers replicate data to their N successors on the ring
- Replication is adjusted whenever a server is added or removed from the ring
- Servers use heartbeat protocol to determine availability
 - Periodic messages exchanged between ring neighbors
 - When a server does not answer heartbeat request in a defined time span, the server is considered gone
- DynamoDB allows reads and writes on every replica!

Data Versioning in Amazon DynamoDB

- DynamoDB follows „always write“ paradigm
 - Write operation allowed on every replica
 - put-() operation returns after one replica has been written
 - ◆ Lazy update of replicas in the background
- Result: Different version of a data item may exist
 - DynamoDB treats each version of the data item as an immutable object
 - Vector clocks are used to reconcile different versions
 - When system cannot reconcile different versions, the versions are presented to client for reconciliation

Overview

- Intro
- Scaling Stateless Components
 - Kubernetes Auto-Scaling
- Partitioning
 - Amazon DynamoDB
- **Replication and Consistency**
- CAP Theorem

Replication

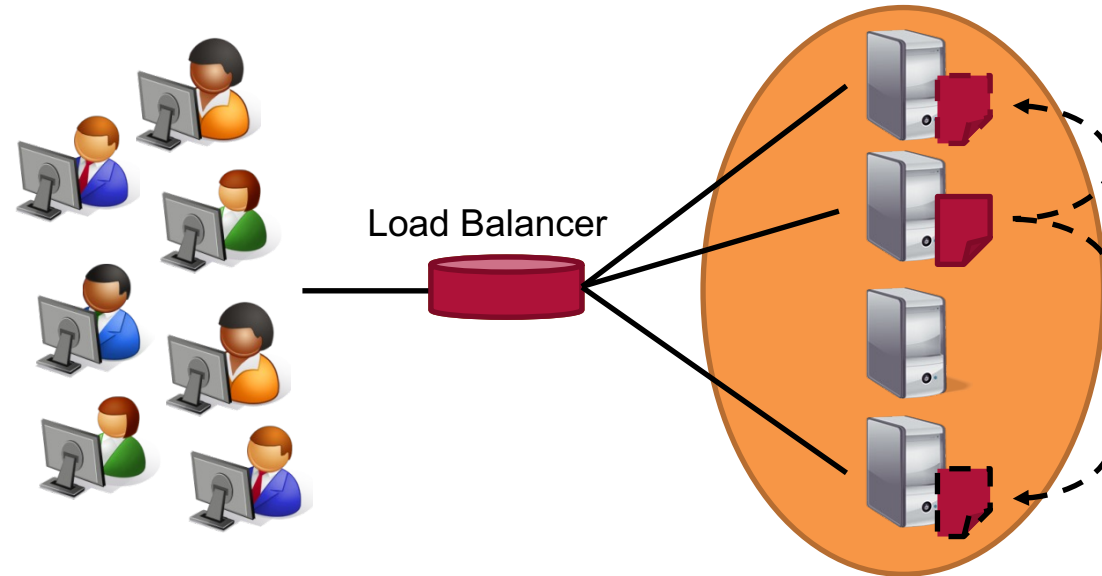
- Partitioning helps with scalability but not availability
 - Data is only stored in one location
 - If a host goes down, its partitions are gone
- Replication: Copies of the data on different machines
 - Where to place the copies?
 - Who creates the copies?
 - What happens when the data changes?
 - How do deal with inconsistencies among the copies?
- Replication can improve both scalability and availability

Where to Place the Replicas?

- Within a cloud data center, replica placement is often done with respect to the network hierarchy
 - One replica on another machine
 - ◆ Guards against individual node failures
 - One replica on another rack
 - ◆ Guards against outages of the rack switch
- On a global scale, replicas are often distributed taking the locations of clients into account
 - Chosen to keep network transfers local
 - Typical with e.g. content distribution networks (CDNs)

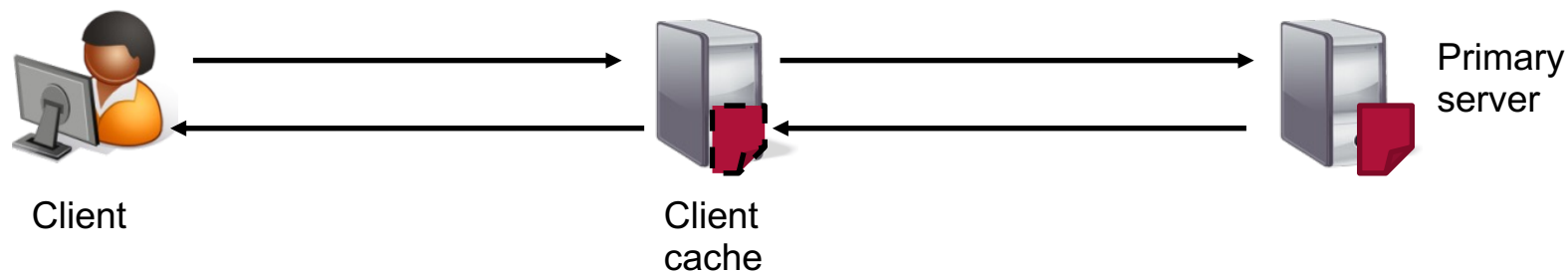
Who Creates the Copies? (1/2)

- Server-initiated replication
 - Mainly used to reduce server load
 - Server decides among a set of replica servers
 - Copies are created by a server when popularity of data item increases



Who Creates the Copies? (2/2)

- Client-initiated replication (== client caches)
 - Replica is created as result of a client's interaction
 - Server has no control of cached copy anymore
 - ◆ Stale replicas handled by expiration date
 - Traditional examples: Web proxies



Push vs. Pull (1/2)

- Push-based updates (server-based protocols)
 - Server initiates replica updates to replica servers
 - Used when high degree of consistency is needed
 - Only usable when all replicas are known
- Pull-based updates (client-based protocols)
 - Clients request updates from server
 - Often used by client caches

Push vs. Pull (2/2)

Issue	Push-based	Pull-based
State of server (about replicas)	Server must keep a list of all replicas and caches	None
Messages sent	Messages are sent only when data changes	Caches poll server, even without updated values
Response time at client	None (except when only invalidation has been transmitted)	Time it takes to fetch the modifications

What Happens When the Data Changes?

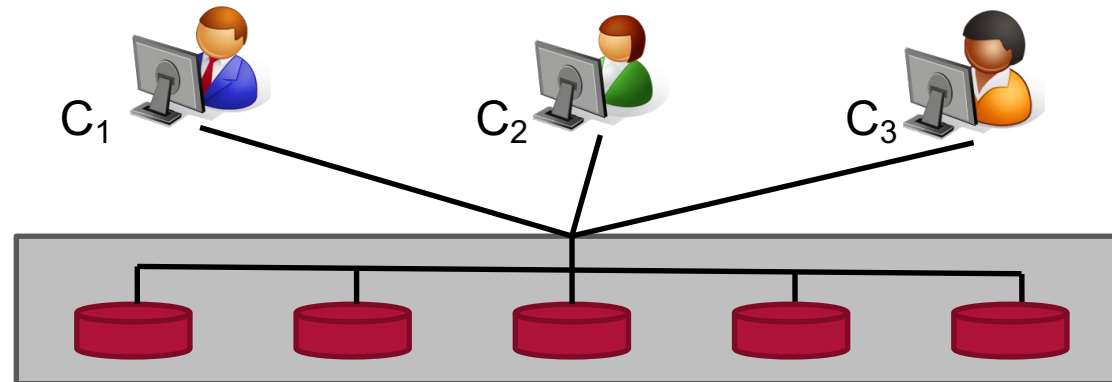
1. Invalidation protocols
 - Inform replica servers that their replica is now invalid
 - Good when many updates and few reads
2. Transferring the modified data among the servers
 - Each server immediately receives latest version
 - Good when few updates and many reads
3. Don't send modified data, but modification commands
 - Good when commands substantially smaller than data
 - Assumes that servers are able to apply commands
 - Beneficial when network bandwidth is the scarce resource

How to Deal with Inconsistencies?

- When data lives in different locations, inconsistencies can occur, for example:
 - Client 1 updates data item X at replica server R1, Client 2 reads old value of data item X at R2
 - Client 1 updates X at R1, Client 2 updates X at R2
 - ◆ Which one is the correct value?
- Different levels of consistency can be enforced
 - Rule of thumb: the higher the consistency level (CL), the lower the performance
 - However, there must be a clear understanding what CL an application can expect from a data store

Consistency Models

- Assume a data store with replication
 - Clients can read/write data stored
 - Clients cannot influence which copy they access
 - Updates are propagated among replicas in data store



- Consistency model: Contract between client and data store → what version of the data can a client expect?

Two Views On Consistency

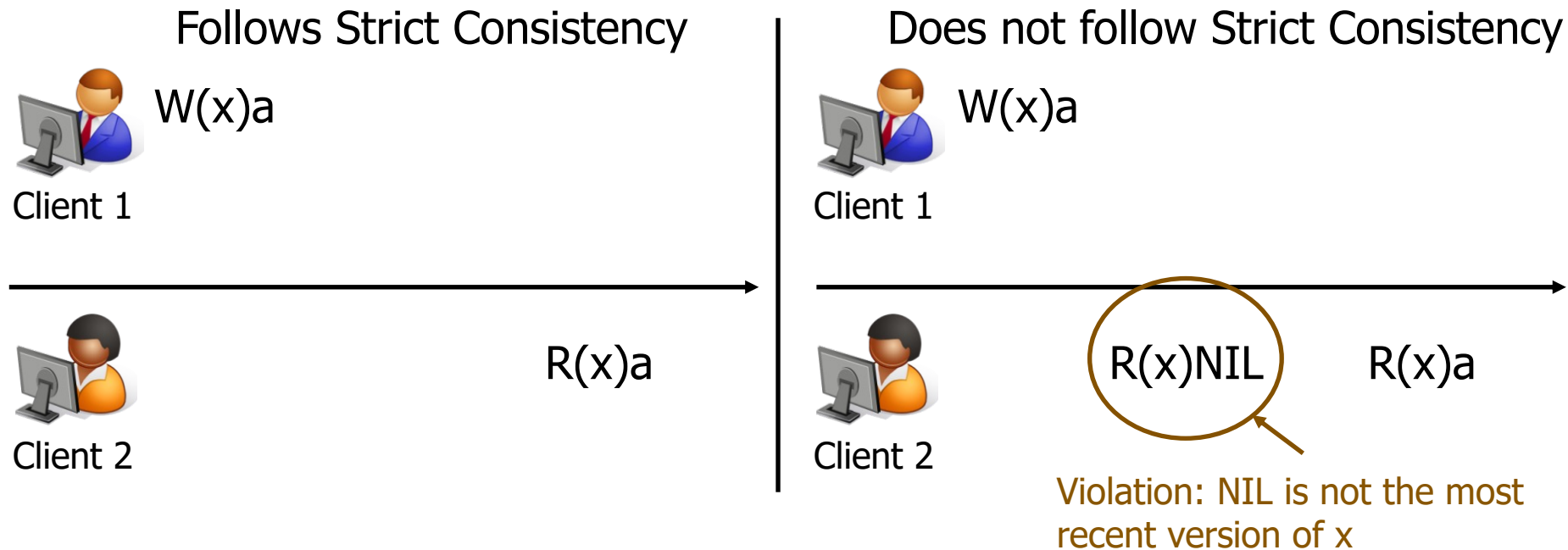
- Data-centric consistency models
 - Talk about consistency from a global perspective
 - Provides guarantees how a sequence of read/write operations are perceived by *multiple clients*
- Client-centric consistency models
 - Talk about consistency from a client's perspective
 - Provides guarantees how a *single client* perceives the state of a replicated data item

Data-Centric Consistency Models

- Strong consistency models
 - Operations on shared data are synchronized
 - ◆ Strict consistency (related to time)
 - ◆ Sequential consistency (what we are used to)
 - ◆ Causal consistency (maintains only causal relations)
 - ◆ ...
- Weak consistency models
 - Synchronization only when data is locked/unlocked
 - ◆ General weak consistency
 - ◆ Release consistency
 - ◆ Entry consistency
 - ◆ ...

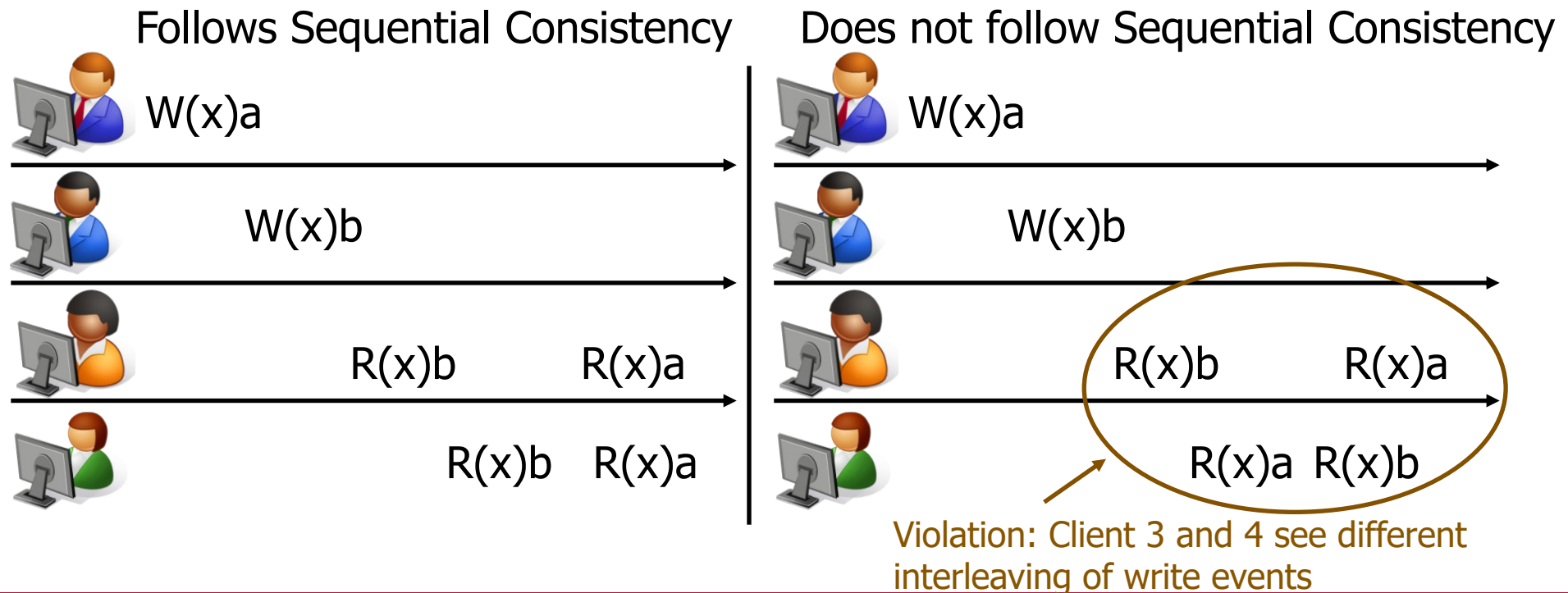
Strict Consistency

- Guarantee to clients: Any read to a shared data item x returns the value stored by the most recent write operation on x
- Model only of theoretical interest in distributed systems



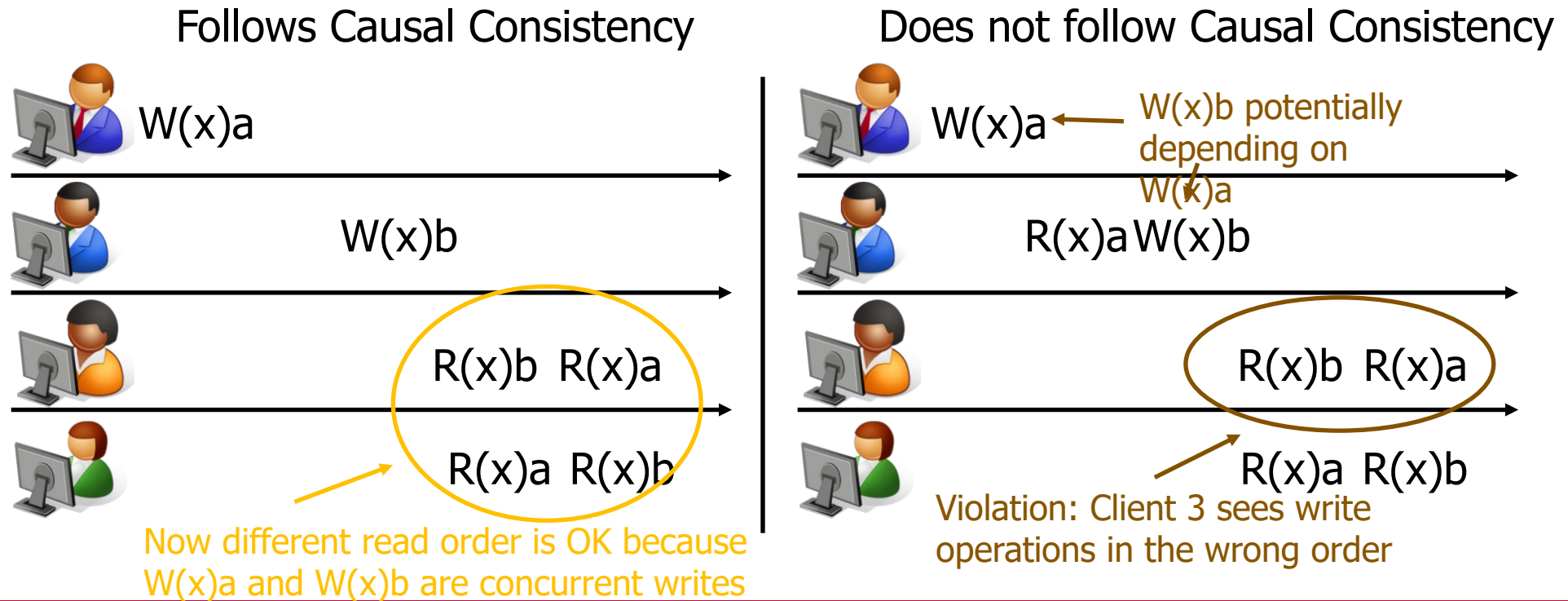
Sequential Consistency

- Guarantee to clients: The result of any execution is the same as if the (read and write) operations of all processes were executed in *some sequential order* and the operations of each individual process appear in this sequence



Causal Consistency

- Guarantee to clients: Writes that are potentially causally related must be seen by all processes in the same order. Independent writes may be seen in a different order by different clients.



Client-Centric Consistency Models

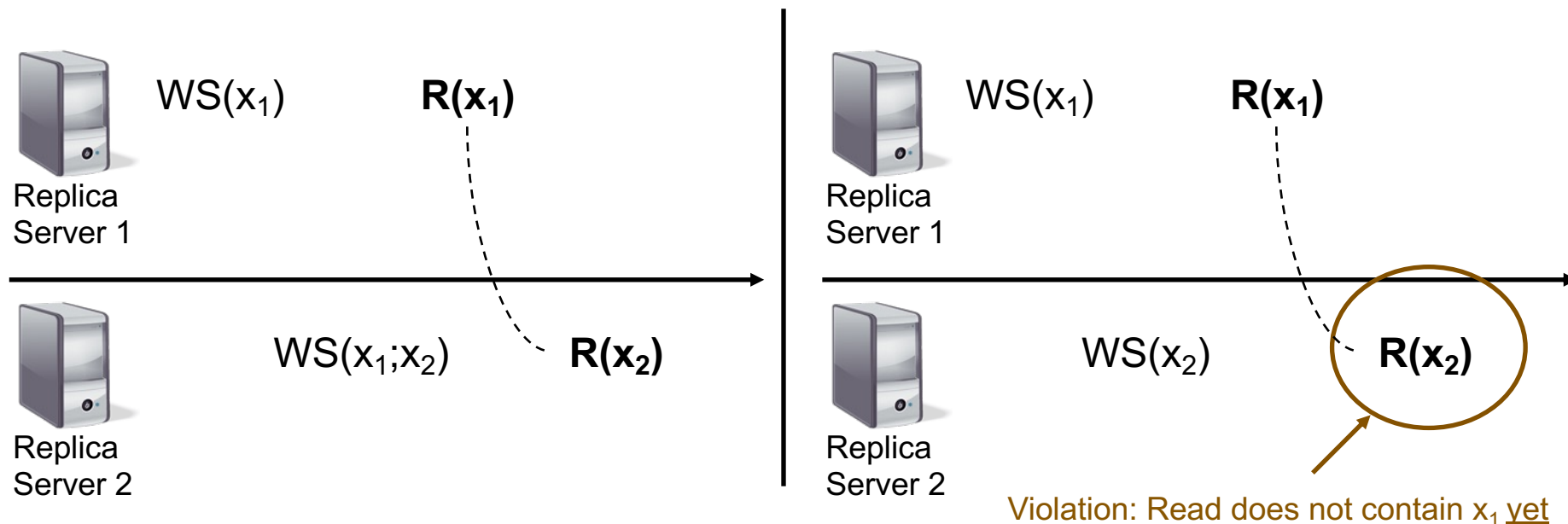
1. Eventual Consistency
2. Monotonic Reads
3. Monotonic Writes
4. Read Your Writes
5. Writes Follow Reads

Eventual Consistency

- Guarantee for the client: All replicas will eventually reach the most recent state
- Apart from that there is no guarantee
 - Client can read old data
 - Client can lose its updates
- Model enjoys popularity in the cloud world as it can be implemented easily and allows high scalability

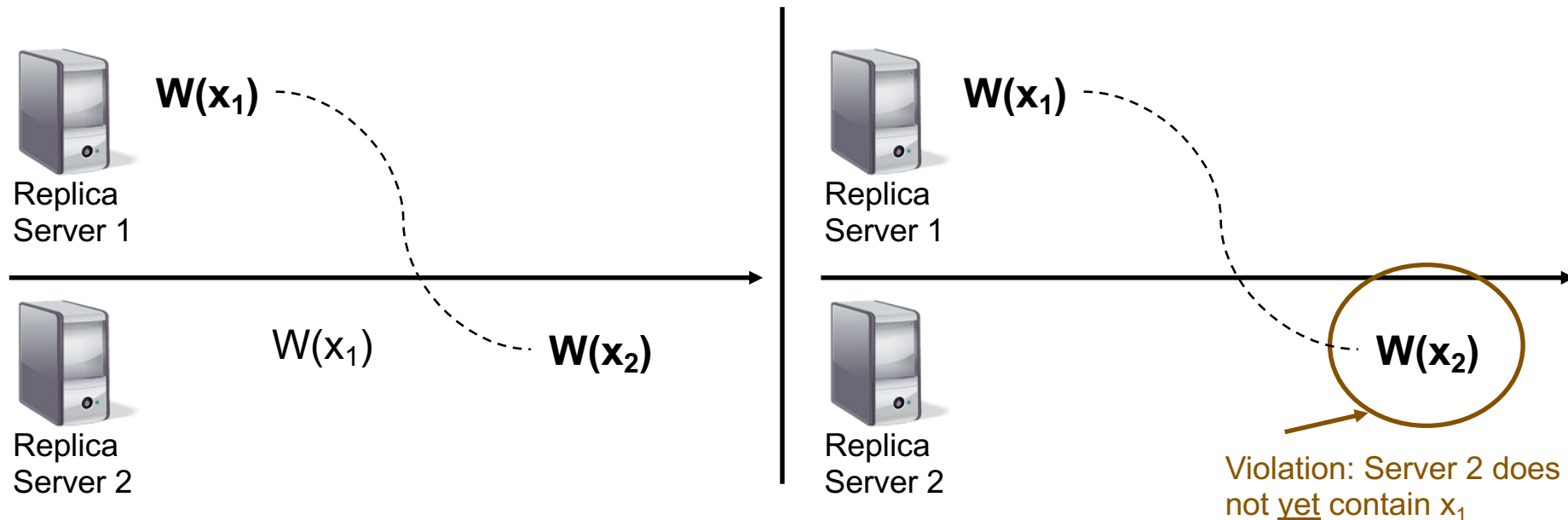
Monotonic Reads

- Guarantee for the client: A read operation by one client is made only at a server containing all writes that were seen by previous reads



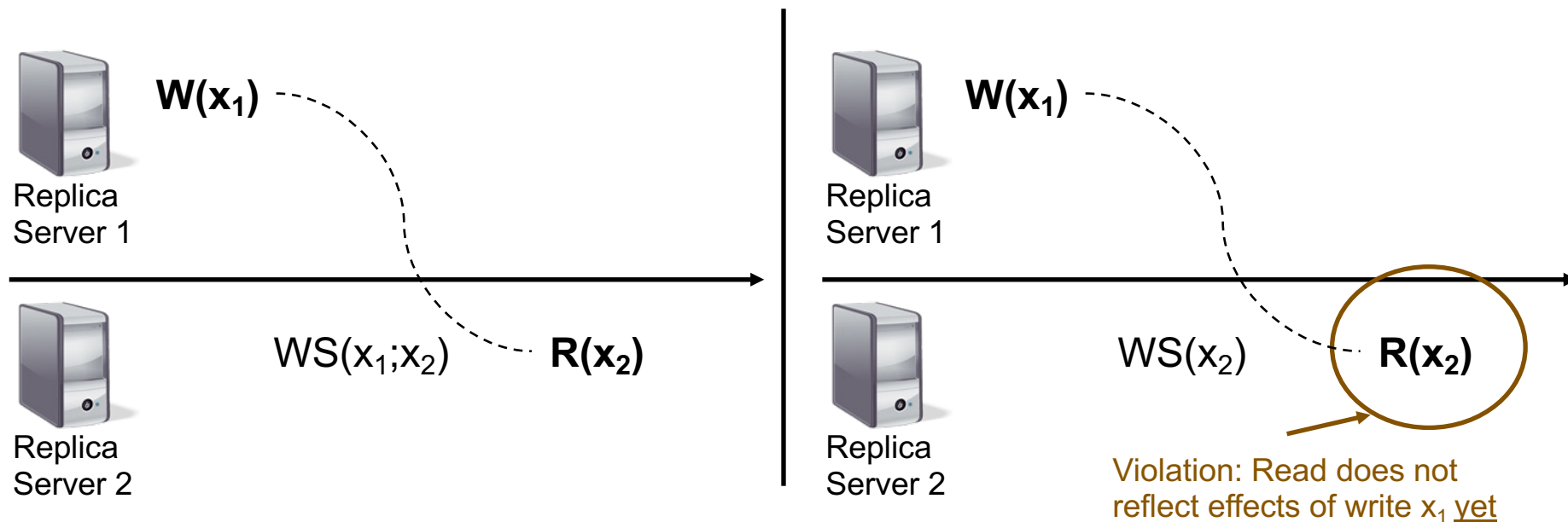
Monotonic Writes

- Guarantee for the client: A write operation by a client on data item x is completed before successive write operations on x by the *same client*



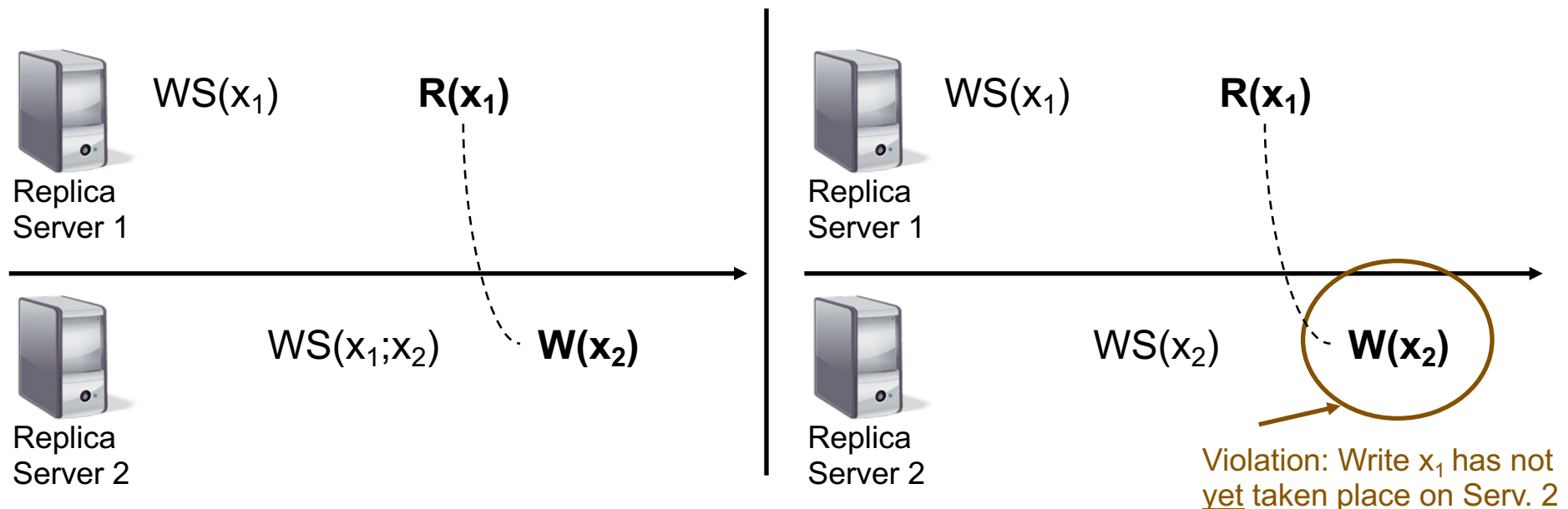
Read Your Writes

- Guarantee for the client: The effect of a write operation by a client on data item x will always be seen by a successive read of x by the *same client*



Writes Follow Reads

- Guarantee for the client: If a read on x precedes a write, then that write is performed after all writes that preceded the read



Summary Client-Centric Consistency Models

- Monotonic-read consistency
 - If a process reads x , any future read on x by the process will return the same or a more recent value
- Monotonic-write consistency
 - A write by a process on x is completed before any future write operations on x by the same process
- Read your writes
 - A write by a process on x will be seen by a future read operation on x by the same process
- Writes follow reads
 - A write by a process on x after a read on x takes place on the same or more recent value of x that was read

General Remark on Consistency

- In general, the stricter the consistency model, the more it impedes the scalability of a system
 - More consistency requires more synchronization
 - While the data is synchronized, some client requests may be answered
- Databases of the 80s and 90s put strong emphasis on consistency, lived with limited scalability/availability
- Today's cloud databases often sacrifice consistency in favor of scalability and availability

Overview

- Intro
- Scaling Stateless Components
 - Kubernetes Auto-Scaling
- Partitioning
 - Amazon DynamoDB
- Replication and Consistency
- **CAP Theorem**

Brewer's CAP Theorem^[2]

- In a distributed system, it is impossible to provide all three of the following guarantees at the same time:
 - Consistency: Write to one node, read from another node will return something no older than what was written
 - Availability: Non-failing node will send proper response (no error or timeout)
 - Partition tolerance: Keep promise of either consistency or availability in case of network partition

- Proof by Seth Gilbert and Nancy Lynch^[3]

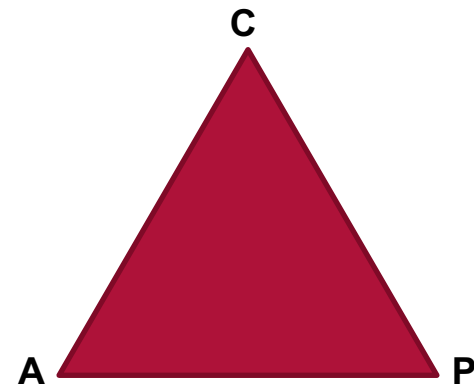


Illustration of CAP Theorem (1/2)

- Illustration of CA (consistency and availability)
 - Example: Replicated DBMS
 - Provided all servers can communicate, it is possible to achieve consistency and availability
- Illustration of PC (partition tolerance and consistency)
 - Example: Pessimistic locking of distributed databases
 - System can provide consistency even if some nodes are temporarily unreachable
 - However, some requests may not be answered due to unreachable nodes (violates availability property)

Illustration of CAP Theorem (2/2)

- Illustration of AP (availability and partition tolerance)
 - Example: DNS system
 - System continues to work even when some DNS servers are offline (availability)
 - System tolerates network outages (partition tolerance)
 - DNS makes extensive use of caching to achieve partition tolerance
 - ◆ Non-authoritative DNS server may answer request with old data (violates consistency)

Summary

- Scaling-out to more virtual resources: scalability and fault tolerance
 - Load balancing for replicated stateless components
 - Load balancing *and* data consistency models for replicated stateful components
- The higher the consistency level, the less scalable replicated services are
- System components fail eventually, decide for either availability or consistency as a system's designer

Literature and References

- Literature

- A.S. Tannenbaum, M. Van Steen: “Distributed Systems: Principles and Paradigms”, Prentice Hall, 2016, Chapter 7
- M. Kleppmann, “Designing Data-Intensive Applications”, 2017, Chapter 5 and 6

- References

- [2] Eric. A. Brewer: “Towards Robust Distributed Systems”, PODC Keynote 2004, <http://www.cs.berkeley.edu/~brewer/cs262b-2004/PODC-keynote.pdf>
- [3] S. Gilbert, N. Lynch: “Brewer’s Conjecture and the Feasibility of Consistent, Available, Partition-Tolerant Web Services”, ACM SIGACT News, 33 (2), 2002
- [4] G. DeCandia, D. Hastorun, M. Jampani, G. Kakulapati, A. Lakshman, A. Pilchin, S. Sivasubramanian, P. Vosshall, W. Vogels: “Dynamo: Amazon’s Highly Available Key-Value Store”, in Proc. of the 21st ACM SIGOPS Symposium on Operating Systems Principles, 2007
- [6] D. Karger, E. Lehman, T. Leighton, R. Panigrahy, M. Levine, D. Lewin: “Consistent Hashing and Random Trees: Distributed Caching Protocols for Relieving Hot Spots on the World Wide Web”, in Proc. of the 29th ACM Symposium on Theory of Computing, 1997
- [7] S. Nakamoto, “Bitcoin: A Peer-to-Peer Electronic Cash System”, White paper, 2008